

BOOK REVIEW

# Influence, New and Expanded review

A book by Robert B. Cialdini

## The limits of getting a yes

Evidence assessed through 15 September 2026 · 5,328 words · 18 min read

OVERALL

61%

An editorial rubric score. It is not the percentage of the book that is true.

|                     |     |
|---------------------|-----|
| Scientific Accuracy | 50% |
| Reference Accuracy  | 67% |
| Practical Value     | 67% |

Robert Cialdini goes into an electronics store and comes out with a television he had not intended to buy.

The set has a good rating and an attractive sale price. A salesman tells him it is the last one. Another customer has called about it and might come in that afternoon. Cialdini knows exactly what is happening. He has spent his professional life studying persuasion. He recognizes the scarcity appeal as the salesman makes it. Twenty minutes later, he is wheeling the television out of the store.<sup>1</sup>

This is Cialdini's own account, near the end of *Influence*. He does not offer it as an embarrassing confession that his life's work has failed. Whether he was duped, he says, depends on whether the salesman was telling the truth. He returns the next morning, finds no replacement television on display, and takes that as confirmation. He writes a favorable review of the store and the salesman.

There is something admirable about the story. Cialdini does not pretend that knowing psychology places him above the people he studies. He also recognizes that a salesperson can supply useful information. Learning that a worthwhile opportunity is about to disappear is not necessarily being manipulated.

But his test leaves an important question unanswered: was buying the television a good use of his money?

The scarcity could have been genuine and the purchase unnecessary. The television could also have been an excellent choice. The account does not settle that question. Checking whether the salesman lied is not the same as checking whether the purchase served the buyer.

Cialdini actually gives readers the stronger test earlier in the book: ask what you want from the scarce item. His cookie example makes the distinction concrete. A cookie can become more desirable because there are only a few left without becoming any better to eat. Likewise, a television can become more urgent to buy without becoming more useful to its buyer. His closing verdict on the television does not apply that earlier test. It makes the seller's honesty decisive.<sup>2</sup>

That leaves three questions to keep separate: did the tactic help make the sale, why did it work, and was the purchase good for the buyer? *Influence* often supplies useful evidence for the first. Its answers to the other two need more scrutiny.

My verdict is that the book is worth reading, selectively. There are real experiments and useful ideas here, and Cialdini is often more careful than the slogans associated with his book. But he sometimes describes a comparison incorrectly, tells us why a tactic worked when the study cannot answer that question, or recommends more than the evidence supports. You can learn from the book without accepting every explanation that comes with the lesson.

## What Cialdini is really arguing

Cialdini's starting point is sensible: we cannot investigate every decision from scratch. We often use a manageable piece of information—a recommendation, a favor, a sign that something is popular—to decide how to respond. He groups these influences into seven principles:<sup>3</sup>

- **Reciprocation:** a favor can make you feel that you owe something in return.
- **Liking:** a request can become more appealing when you like the person making it.
- **Social proof:** other people's choices can guide your own, especially when you are unsure what to do.
- **Authority:** people may follow someone who has relevant expertise—or someone who merely appears to have it.
- **Scarcity:** something can become more desirable when access to it is limited.
- **Commitment and consistency:** an earlier choice or promise can make you feel pressure to follow through.
- **Unity:** someone you regard as “one of us” can have influence that goes beyond being likable.

Unity is the addition in this expanded edition. Together, the principles offer ways to understand how a request gains force. The same information can help you make a good decision or be used to steer you toward a bad one.

Cialdini does not claim that everyone can be made to do anything. He explicitly says that none of the principles works on all people in all situations. He also distinguishes the tasks a persuader might face: building a relationship, reducing someone's uncertainty, and getting them to act are different problems. Money and other direct benefits sit outside his seven principles because he takes their importance for granted.

The qualifications become more substantial when you read the notes. Small favors lose their obligating force with time. Gratitude is not identical to indebtedness. Close relationships should not operate as running accounts of equivalent favors. An exaggerated opening demand can backfire. Combining several tactics can make a message seem heavy-handed. Contact between groups can help or hurt, depending on what happens during that contact. Unity can encourage cooperation and protect misconduct.<sup>3</sup>

He also makes a distinction worth keeping between finding genuine common ground and inventing similarities to make someone like you. His ethical preference is not merely decorative: the book repeatedly criticizes counterfeit expertise, fake popularity, and fabricated scarcity.

Read this way, *Influence* aims to help you make better requests and recognize pressure without becoming suspicious of every kindness or recommendation. That is a worthwhile aim. The question is how much of the book's explanation and protective advice its evidence can support.

Its engaging stories help the reader notice things that might otherwise pass unnoticed: the favor before the request, the smaller offer after the unreasonable one, the change in a deal after a commitment. But recognizing a pattern is the beginning of an explanation. A sales tactic can work for reasons other than the one the observer assigns to it. A story about a successful practitioner cannot tell us how many similar attempts failed.

## **Influence, explanation, and protection**

Start with the claim that most techniques for getting people to agree draw their power from these seven principles. Cialdini also argues that choosing the right principle and using it well can make a request more successful, sometimes substantially, in business and everyday life.<sup>4</sup>

There is substantial evidence for particular uses: experiments have changed purchases, attendance, cooperation, and energy use. Support appears across all seven groups of techniques, though it is stronger and more widely tested in some than others. You do not need seven separate scientific laws to find this useful. Still, a collection of successful examples cannot tell us whether the principles explain *most* ways of getting agreement, or how reliably a technique will work somewhere new. That is why I judge the broader claim only partly supported.

The second claim is about how people decide. Cialdini argues that much agreement comes from responding automatically to a prominent feature of the situation—such as popularity or an expert's title—instead of carefully examining the request. He says these shortcuts usually point toward a beneficial choice, and that we rely on them more when time, energy, motivation, or the ability to think carefully is limited.<sup>4</sup>

Consider a customer who chooses a dish after learning it is popular. The customer might accept the recommendation without much thought. Or they might think it through: a dish that many other diners order seems a reasonable bet. More orders alone cannot tell us which happened. And even if the information increased sales, we still need to know whether customers enjoyed the meal. The book includes evidence about when people examine arguments more carefully, but many of its examples show only what people did, not the thought process that produced it.

The third claim concerns protection: learn to recognize these pressures and use the book's defenses, and you can resist exploitation while still accepting useful favors, expertise, and recommendations.<sup>4</sup>

Some of those defenses have been tested directly. Others sound reasonable but have little direct evidence showing that they protect people. Cialdini himself says that recognizing scarcity pressure may not be enough to resist it. We need to examine each proposal: would treating a sales gift as a sales tactic help, would checking an expert's relevant qualifications help, and would separating a likable person from their offer help?

## What works

In one experiment, a man gave a participant an unsolicited Coke and later asked him to buy raffle tickets. Participants who received the drink bought almost twice as many tickets as those who did not. The original report supports Cialdini's description of the effect.<sup>5</sup>

The study involved 77 male students, so it cannot tell us how everyone responds to a gift. But it measured an actual purchase, not just a promise to buy. The gift changed what the participants did. Whether they bought more because they felt obliged, or for some other reason, is a further question.

A restaurant study gives us a chance to ask about both the sale and the customer's experience. Researchers tested information identifying popular dishes and included a comparison intended to separate the effect of popularity from simply drawing attention to an item. Their report describes a 13–20% increase in orders for the five highlighted dishes, consistent with Cialdini's account, and modest evidence that diners enjoyed their meals more.<sup>6</sup>

That last outcome matters. It offers some evidence on the customer's experience as well as the restaurant's sales. In this setting, social information could help diners choose something they enjoyed. That is exactly the kind of recipient benefit the book's larger argument needs.

Another study examined repeated reports comparing households' energy use with that of other households. Allcott and Rogers asked what happened when those reports stopped after two years. Some of the energy savings lasted, but the effect shrank by roughly 10–20% a year.<sup>7</sup>

That percentage describes the *saving*, not the household's total electricity use. For a simple illustration, suppose the reports had produced a saving of 100 units of electricity. Losing 10% of that effect would leave a saving of 90 units. It would not mean the household's electricity use had jumped by 10%. The distinction matters if you are deciding whether a benefit is likely to last.

These studies give the book practical value. They also show why the details matter. Popularity information might help you choose from a restaurant menu without telling you enough to buy an unfamiliar, expensive service. A benefit that follows two years of reports does not establish what a single message would achieve. When you copy a technique, ask which parts of the tested situation your own situation actually shares.

## Getting agreement is not the same as getting action

The door-in-the-face technique begins with a large request that will probably be refused. The persuader then asks for the smaller favor they wanted all along. That retreat can make the second request look like a concession worth returning. In a classic experiment, agreement to chaperone a zoo trip rose from about 17% to 50% after a larger request was refused. Cialdini also discusses whether people carry out their agreements and cooperate again later.<sup>8</sup>

The agreement finding has later support. Researchers repeated another experiment from the same 1975 paper, specifying their plan before collecting the data—a safeguard called preregistration. In this later study, available online in October 2020, about 51% agreed after the persuader retreated from a larger request, compared with 38% who received only the smaller request.<sup>8</sup>

To find out about follow-through, we need studies that measure it. A 2012 review combining results from multiple studies found a small average benefit for agreement; its estimate for actual behavior was smaller and too uncertain to establish a benefit. An updated review in December 2025 found modest average benefits for both. The estimate for behavior was again smaller, but the two estimates alone do not establish that the true effects differ. There is support for the technique at both stages. Saying yes and showing up still have to be measured separately.<sup>8</sup>

Imagine organizing volunteers for an event. An extra yes is valuable only insofar as it produces useful participation. It can even create a planning problem if the organizer mistakes reluctant agreement for reliable attendance. The relevant outcome is not how persuasive the conversation felt. It is who turns up, what they do, and what the interaction costs the relationship.

Cialdini deserves credit for raising the follow-through question himself. The later evidence strengthens the case for his technique.

A second example shows how easily those numbers get mixed up. In the low-ball experiment, people were asked to join a study with an inconvenient early start. One group learned the start time before agreeing; the other learned it only afterward. Cialdini reports agreement rates of 24% and 56%. But 24% was the first group's *attendance* rate, while 56% was the second group's *agreement* rate.<sup>9</sup>

Compare agreement with agreement and the figures are 31% versus 56%. Compare attendance with attendance and they are 24% versus 53%. Both comparisons still favor delaying disclosure of the start time. The effect survives the correction; the book has described it incorrectly.

The much larger 95% follow-through figure needs its own explanation. That is 18 of the 19 people in the low-ball group who had already agreed, not 95% of everyone asked. If you are planning a study, you need both numbers: how many people agree, and how many of those people attend.

## What the “because” experiment actually tests

The copying-machine experiment is one of the book's clearest examples of what Cialdini calls “click, run”: a familiar signal sets off a response without much thought. Researchers approached people who were about to use a library copier and asked to use the machine first. Adding an apparently empty reason—needing to make copies—worked almost as well as saying they were in a rush. Since everyone at a copier needs to make copies, the empty reason explains nothing about why this person should go first.<sup>10</sup>

When the favor was relatively small, 60% agreed without a reason, 93% with the empty reason, and 94% with the meaningful reason. But when the favor was larger, only 24% agreed in both the no-reason and empty-reason conditions. “Small” and “large” referred to the amount of copying requested relative to the other person’s task, not simply to a fixed number of pages.<sup>10</sup>

The empty explanation helped in one situation and produced the same agreement rate as no reason in the other. Cialdini acknowledges that people do not always respond mechanically, but he does not put this comparison beside the memorable finding. Instead, he emphasizes *because* as the trigger. The experiment, however, compared differently worded requests. It did not change only the word *because* while leaving everything else the same. We therefore cannot tell whether that one word deserves the credit.

A later attempt by Valerie Folkes to repeat and extend the study challenged its explanation. Langer and colleagues disputed that challenge, arguing that the new conditions changed the test and sometimes left little room for agreement rates to rise further. The researchers disagreed about when people pay attention to reasons. That dispute should make us cautious about the explanation, without pretending the original result never happened.<sup>11</sup>

The practical question becomes more interesting than whether to insert a magic word. What are you asking the other person to give up? A phrase that accompanies agreement to a small interruption may do nothing when the request is more burdensome.

Elsewhere, Cialdini uses an experiment by Petty, Cacioppo, and Goldman to show that people weigh expertise and argument quality differently depending on how personally involved they are in the issue.<sup>12</sup> That supports his view that people do not always process a message in the same way. The experiment measured attitudes, though. It cannot tell us that the purchases, favors, and acts of obedience elsewhere in the book all came from the same automatic process.

## How much can a reassuring phrase do?

In the scarcity chapter, Cialdini discusses a simple way to reduce resistance: tell people they are free to refuse. Acknowledging their choice might make them more willing to cooperate. The book cites a meta-analysis—a review that combines the numerical results of multiple studies—rather than relying only on a colorful encounter.<sup>13</sup>

A later review gives reason to be less confident. Fillon and colleagues’ 2023 meta-analysis covered 52 experiments involving 19,528 participants. Taken together, the studies favored the phrase. But when the researchers examined only the seven studies whose methods they judged least likely to distort the results, the estimated benefit became much smaller and too uncertain to establish a benefit.<sup>13</sup>

This is not proof that the phrase does nothing. The stricter analysis still leaves room for a useful effect, the quality judgments are open to debate, and a broader group of studies with some quality concerns retained a positive result. What we have is less reason to promise a dependable boost in agreement.

Giving someone a real choice can still be the right thing to do, even if it does not help you get your way. But “you are free to refuse” offers little reassurance if refusal brings a penalty. Permission in the sentence is not necessarily freedom in the situation.

## Commitments can help, but the comparison matters

Can a public promise help people follow through on something worth doing? Cialdini describes clinicians displaying signed commitments to avoid prescribing antibiotics when they were inappropriate. The study involved fourteen clinicians and checked what they prescribed. That gives us more to work with than a story about someone determined to keep a promise.<sup>14</sup>

The comparison is important. The book says clinicians with signed pledges were compared with clinicians displaying ordinary information posters. In fact, the comparison group had no poster. The study therefore tested the signed-poster approach against usual care, not whether adding a personal promise worked better than putting up information alone. The benefit was measured during a twelve-week poster intervention at the end of the study year. The study did not test whether it persisted after the posters were removed.<sup>14</sup>

The main result was favorable. After the researchers' statistical adjustments, inappropriate prescribing fell from 43.5% to 33.7% in the poster group, a drop of 9.8 percentage points. It rose from 42.8% to 52.7% in the comparison group, an increase of 9.9 points. The poster group moved 9.8 points down while the comparison group moved 9.9 points up, making their changes 19.7 points apart. To find a difference in percentage points, subtract one rate from the other: 43.5 minus 33.7 is 9.8. It is not a 9.8% reduction relative to the starting rate.

The book also reports a 27% reduction, but that figure is harder to reconcile. Part of the problem lies in the original paper: its opening summary and main results table reverse the groups' starting-rate labels. Neither version gives a conventional 27% reduction from the starting rate. The book's separate 21% increase in the comparison group can, however, be traced to the opening summary. It would be unfair to blame Cialdini for every mismatch. Note 14 and the appendix show the figures and calculation limits.

The useful result is a reduction in inappropriate prescribing while the posters were displayed. The study did not measure an improvement in patients' health. That distinction matters when deciding what benefit you can reasonably promise from copying the intervention.

Would the same idea work elsewhere? A 2020 English trial randomly assigned 196 practice units and did not find a benefit from posters on its main measure, overall antibiotic dispensing. That trial delivered the intervention differently and counted all antibiotic dispensing, rather than focusing on inappropriate prescribing as the US study did. It was not an identical repeat of the earlier experiment. It does show why distributing posters in a new setting is not enough to assume the original result will follow.<sup>15</sup>

Cialdini is more careful about the explanation in his Iowa energy-conservation story. He suggests that people came to see themselves as conserving energy and developed their own reasons to continue. He presents that as the explanation he favors, rather than an established fact. His notes also discuss findings in which making goals public did not help.<sup>16</sup> That is the distinction to keep: a commitment can change what someone does without proving a permanent change in who they think they are.

Before making a public pledge, ask whether the goal is worth pursuing and whether the promised action is something you can do. Keeping a promise is not automatically a good outcome. Remaining committed to a bad purchase or an obsolete plan is one of the problems the book itself wants readers to recognize.

## The promise and limits of unity

Unity is different from liking. You can like a helpful salesperson without feeling that you and the salesperson belong to the same group. Cialdini's unity principle concerns the extra pull of someone you experience as *one of us*. His chapter explores families, teams, communities, shared activities, and people getting closer by disclosing things about themselves to each other. It also asks whether those connections can cross existing group boundaries.<sup>17</sup>

There is experimental support for part of that ambition. Salma Mousa randomly assigned displaced Iraqi Christian soccer teams to include Muslim players or remain all-Christian. Players on mixed teams showed improvements in some behaviors toward Muslim peers, including reporting that they trained with Muslims at a six-month follow-up. But these gains did not extend much to life away from soccer or to broader attitudes.<sup>18</sup> The study gives us evidence of useful cooperation in and around a shared activity. It gives much less reason to expect that cooperation to change someone's views of an entire group.

Cialdini also warns that a strong group identity can protect misconduct, that people can cooperate toward bad goals, and that efforts to reduce prejudice may fail or fade. His notes cite Mousa as part of a general caution about contact between groups. They do not spell out the soccer study's specific limit: improvement among teammates did little to change attitudes or behavior more broadly.

The difficulty comes when the chapter moves from what people do together to explanations involving genetic similarity, brain activity, and a shared sense of identity. These are different things to measure. Evidence that teammates cooperate does not, by itself, tell us what happened in their brains or whether genetic similarity explains their behavior. The appendix examines those separate claims. A medical example in the chapter's notes shows why keeping an observation separate from its explanation can matter.

Cialdini cites a 2020 study by Greenwood and colleagues as evidence that newborn mortality declines when the physician and newborn have the same race. The favorable finding concerned Black newborns in Florida hospital records, not every possible racial pairing. Babies were not assigned to

doctors at random. That leaves a problem: a difference in outcomes could reflect differences in the patients the doctors treated, as well as differences in their care.<sup>19</sup>

The comparison needs care too. One important estimate asks whether the mortality gap between Black and White newborns differs according to the physician's race. In other words, it compares *two gaps*: the Black–White gap among patients of White physicians and the Black–White gap among patients of Black physicians. That is not simply a comparison of Black babies treated by Black doctors with Black babies treated by White doctors.

In 2024, Borjas and VerBruggen reanalyzed the same records. When they replaced the earlier adjustment for common diagnoses with measures accounting for very low birthweight, the estimated difference between those two racial gaps became much smaller and was no longer statistically significant.<sup>19</sup> The changed analysis did not provide clear evidence that the two gaps differed. It did not prove that the difference was exactly zero.

Why does that change matter? Imagine comparing two doctors when one treats more babies who arrive at much greater risk. You would not want to attribute the entire difference in deaths to the doctors' care. Accounting for patients' starting risks is an attempt to make the comparison fairer. The later paper shows that this finding depends substantially on how those risks are handled. The original researchers had already warned that unmeasured differences might affect their estimates.

The later analysis uses the same population, so it is not a fresh test in a different set of hospitals or patients. Neither paper settles whether racial bias or shared identity affects care in general. Together, they leave this study an insecure basis for claiming that racial matching itself saves newborns—or that unity is the reason. A place in a persuasion chapter is not enough to turn that interpretation into advice about choosing a doctor.

## Learning to resist

The book's advice about authority has a more direct test. Can people learn to distinguish someone who knows the relevant subject from someone who merely looks impressive? In research with Sagarin and colleagues, training helped participants make that distinction in advertisements, including ads they had not seen during training. Showing participants that they too could be influenced by a misleading authority appeal could strengthen the training's effect.<sup>20</sup>

This is evidence for one of the book's defenses, with limits. Participants became better at evaluating ads. A delayed test followed them for one to four days, and many of the original participants did not return. The study did not establish that the training prevented unwanted purchases or protected people for months. Those would be useful next questions, rather than reasons to ignore the skill participants did learn.

The fake-review box is headed "Here's How to Spot Fake Online Reviews with 90 Percent Accuracy, according to Science." That invites a claim about the reader's skill. The text immediately identifies a computer program, offers linguistic clues, and acknowledges that those clues alone will not make

someone a master detector.<sup>21</sup>

Cialdini therefore does tell us that the 90% belongs to a computer program. The problem is the invitation in the headline: *here is how you can do this*. In the underlying paper, three human judges classified reviews as real or fake with about 53–62% accuracy on a subset of 160 reviews. They had not been taught Cialdini's advice, so their scores do not show whether his tips work. The computer's score does not answer that question either.<sup>21</sup>

A headline about a computer detecting fake hotel reviews would fit the study. Turning its result into advice for readers requires another step: teach people the suggested clues and test whether their judgments improve. The program's results can suggest clues to investigate; they cannot substitute for that test.

The reciprocity chapter offers a different kind of defense. Cialdini sometimes says that recognizing a gift as a sales device will free you from its pull. That recognition can give you a reason to say no. It does not show that the feeling of owing someone disappears as soon as you understand what they are doing.<sup>22</sup>

Nor do you have to wait for that feeling to disappear before declining. This is my practical addition, not an effect demonstrated by the book: you can refuse an unwanted request while still feeling indebted. Treating discomfort as something you must eliminate first would make saying no harder than it needs to be.

## **An honest cue can still lead to a poor choice**

The closing chapter says that these tactics become exploitative only when the relevant information is made up. Cialdini wants to preserve useful trust: fake popularity and counterfeit expertise spoil information we could otherwise use. His toothpaste example also asks us to assume that we already want good toothpaste. He has not forgotten that the buyer has a purpose.<sup>2</sup>

A poor purchase alone does not prove that anyone was exploited. My concern is with the claim that exploitation requires a lie. Suppose a manager invokes an employee's genuine promise to help finish a project, pressuring them to cancel care for a sick child to do routine work that could wait. The manager knows the employee is afraid of losing the job and deliberately uses that dependence. The promise is real, but so is the avoidable cost imposed on the employee. I would regard that as exploitation even without a lie. The truth of the reminder does not settle whether using it this way is acceptable.

Cialdini's earlier scarcity advice gives us another question to ask: what do I want from this thing? A truthful statement that an item is nearly sold out does not answer that question. Checking the seller's claim and checking your own reasons for buying are both useful.<sup>2</sup>

The unity chapter raises another workplace question: what should an organization do when a member breaks its code? Cialdini recommends dismissal for a proven major violation or multiple proven minor violations. This is not a rule to fire someone for every mistake. The studies examined for this review support concern about dishonesty and how organizations respond, but they do not directly test that particular dismissal rule.<sup>23</sup>

His insistence on proof, and his distinction between a major violation and repeated minor ones, are useful boundaries. An organization still has to decide how to establish what happened, protect those affected, and respond proportionately. Research about misconduct can inform those decisions without settling them. Cialdini's recommendation may be defensible as judgment; the presence of nearby studies does not turn it into a scientifically established policy.

## How I would use the book

Begin with the decision you are trying to improve. For someone making a request, that means asking whether the proposal deserves acceptance, not merely which principle might obtain it. For someone receiving a request, it means asking whether the offer still makes sense once the urgency, favor, flattery, or affiliation is set beside the actual costs and benefits.

Use the principles to generate questions. Why does this offer suddenly feel attractive? What does the apparent expert know about this specific issue? Is the popularity information relevant to my needs? Am I continuing because the choice remains worthwhile, or because changing my mind feels uncomfortable? These questions develop the book's stronger defensive passages. I would use them to structure a decision, without treating a feeling of recognition as proof that I have diagnosed the situation.

For organizations, check the result you actually want. If you need volunteers, measure useful participation. If a message increases sales, ask what happens after the sale too. If a pledge is meant to change behavior for months, a brief improvement does not yet answer your question. These are my proposed standards for using the research.

Preserve genuine common ground and useful expertise. Treating every kindness or recommendation as a trap can make useful cooperation harder. Cialdini is right about the costs of blanket distrust. The alternative is selective evaluation: take a cue seriously when it supplies relevant evidence, and remain free to reject the proposed action even when the cue is real.

The strongest reason to read *Influence* is that some of its techniques changed what people did, and some of its defenses improved judgment in experiments. Knowing the vocabulary may help you notice pressure, but recognizing a principle is not the same as benefiting from it. The recurring problem is the extra work asked of a successful example: it is made to explain why people acted and tell us what to do elsewhere. Later research adds reasons for caution, but some of those limits were visible in the original studies or in evidence available before this edition.

That leads to the 61% overall score. The book contains useful techniques and meaningful safeguards, and many of its references are easy to trace. Some important study descriptions need correction, however, and its broad explanations and promises of protection go further than the evidence. For ordinary uses, the likely benefits are reasonably proportionate to the effort and costs; applying the advice in new settings or to consequential medical and workplace decisions requires more care. The appendix explains all nine grading decisions, including where a different reasonable judgment would change the score.

Cialdini's television purchase returns us to a question his own scarcity chapter teaches us to ask. Knowing the principle did not answer whether he should buy the set. Finding that the salesman appeared honest did not answer it either. What was absent from that closing verdict was the application of his earlier test: does this choice make sense for what I want?

Read *Influence* to become better at asking that question under social pressure. Do not let a persuasive explanation relieve you of answering it.

---

## Preparation and disclosures

This review was prepared at Jason Hreha's request with AI-assisted research, drafting and editorial critique. The preparation covered the full supplied book text, its notes, and all embedded image placements, with small-text legibility limits. The external audit examined 33 selected dossiers across all seven principles. It is a critical narrative review, not an exhaustive systematic search or a count of every claim in the book. Source access varies and is documented in the appendix. Original participant-level analyses were not rerun. Independent human scientific sign-off and legal review are not documented.

Jason Hreha has confirmed that he has no personal or business relationship with Robert Cialdini, Influence at Work, the Cialdini Institute, or the publisher to disclose. He leads The Behavioral Scientist, wrote the commercially available Real Change, and offers behavioral-science consulting and related services. These are relevant commercial and intellectual interests, including potentially competing explanations and products.

This is an independent house review, not an official Red Pen Reviews review or endorsement. Its scores evaluate evidence and practical advice; they do not assess anyone's honesty or character. The analysis provides general information, not individualized medical, psychological, legal, or financial advice. Factual corrections and relevant additional evidence can be submitted through The Behavioral Scientist contact form. Material corrections made after first publication will be dated and explained. The evidence cutoff is not a publication date.

**Publisher:** The Behavioral Scientist / Haystack Group LLC. **Editorial contact:** Jason Hreha.

**Edition reviewed:** *Influence, New and Expanded: The Psychology of Persuasion*. Robert B. Cialdini. Harper Business / HarperCollins, Digital Edition May 2021. Ebook ISBN 9780062937674; version 03252021. These identifiers come from the supplied copyright text [23.0007]–[23.0017].

## Book and review details

**The Behavioral Scientist · Book review**

**Book:** Robert B. Cialdini, *Influence, New and Expanded* (2021).

**Evidence assessed through:** 15 September 2026.

**Overall:** 61% · **Scientific Accuracy:** 50% · **Reference Accuracy:** 67% · **Practical Value:** 67%.

**Recommendation:** Worth reading selectively for useful persuasion techniques and questions to ask under pressure. Its explanations and defensive assurances need more scrutiny.

The scores summarize an editorial rubric. They are not percentages of true statements. The **technical appendix** explains the evidence, nine inputs, source limits, and reasonable alternative grades. Read it with the **house methodology** and **shared editorial notice**.

---

## Evidence notes

Book locators refer to the supplied 2021 edition's ordered spine and paragraph numbers. They are not printed page citations. Full source access and verification limits appear in the appendix.

1. Cialdini's television account: [14.0034]–[14.0040]. These are his reported events, not independently authenticated transactions. The review does not decide whether the purchase was worthwhile or whether checking the display proved the salesman truthful.
2. Book [14.0030]–[14.0033], [14.0043]; its in-the-market condition [14.0031]; scarcity-use distinction [11.0173]–[11.0174]. The original Worchel, Lee, and Adewole (1975) study was read in full and supports the taste/desirability distinction. The examples of truthful pressure are hypothetical illustrations and editorial judgments.
3. Framework and qualifications: [04.0008], [05.0008]–[05.0013], [07.0097], [07.0106], [07.0115], [08.0075], [08.0130], [17.0006], [17.0030], [17.0041]–[17.0042], [17.0095], [17.0182]. These support the attribution of qualifications to the book; checking every underlying cited study is a separate task.
4. The three fixed assessment propositions were selected after reading the book and before detailed external research. They cover the reach of the seven-principle framework, automatic/adaptive shortcuts, and targeted defenses. Their exact wording is reproduced in the appendix; the shorter descriptions here do not replace it. They are reviewer formulations, not quotations from Cialdini.
5. Regan (1971), Effects of a favor and liking on compliance. Original methods, results, tables, and discussion were inspected. The analyzed sample was 77 male students. Published means were 1.73 tickets after a favor and 0.92 without one;  $1.73/0.92 \approx 1.88$  uses those rounded means after the source's recoding. The author reported that the conclusions survived without recoding two unusually large purchases. Book [07.0022]–[07.0026]. See appendix R01 for the third condition and mechanism limits.
6. Cai, Chen, and Fang (2009), Observational Learning: Evidence from a Randomized Natural Field Experiment. The final primary abstract reports the 13–20% demand increase for the five highlighted dishes, salience comparison, and modest evidence of improved dining satisfaction. Final model tables were not retrieved. Book [09.0006]–[09.0007].
7. Allcott and Rogers (2014), The Short-Run and Long-Run Effects of Behavioral Interventions. The final primary abstract supports persistence and annual decay after two years of reports; selected working-paper pages were also inspected. The decay is a fraction of the intervention effect, not a percentage increase in total electricity consumption. Book [09.0149], [17.0087].
8. Book [07.0114]–[07.0115], [07.0141]–[07.0152]. The original Cialdini et al. (1975) experiments and Genschow et al. close replication were inspected at methods/results scope. The replication targeted Study 3 of the original paper and was online in October 2020; 391 were analyzed, with 51.28% versus 37.88% agreement in the relevant arms. The 2012 synthesis was checked at primary-abstract scope, as was Feeley's 2025 update. The 2012 pooled correlations were  $r=.126$  for verbal agreement and  $r=.052$  for behavior; the latter was statistically nonsignificant. The December 2025 update reports positive estimates of  $r=.14$  across 89 verbal-compliance comparisons and  $r=.08$  across 53 behavioral comparisons. These correlations are not percentage increases, and the two estimates alone do not establish a statistically significant difference between outcomes. These source checks do not independently verify every older follow-through table, study overlap, or moderator analysis.

9. Cialdini, Cacioppo, Bassett, and Miller (1978), Low-ball procedure for producing compliance. Experiment 1 methods/results were inspected: low-ball versus upfront-time agreement was 19/34 versus 9/29; attendance was 18/34 versus 7/29. The book's 95% is 18/19 among those agreeing. Book [12.0212]. A favorable Burger and Caputo (2015) synthesis was also read; it does not settle a unique mechanism.
10. Langer, Blank, and Chanowitz (1978), The mindlessness of ostensibly thoughtful action. Study 1 methods and Table 1 were inspected, including the table image. Small-favor counts were 9/15, 14/15, and 15/16 in the no-reason, redundant-reason, and real-reason conditions. Large-favor counts were 6/25, 6/25, and 10/24. The real-reason contrast for the larger favor was only marginal in the source's test. Book [06.0018]–[06.0020]; flexibility qualification [17.0011].
11. Folkes (1985), Mindlessness or mindfulness, and Langer and colleagues' reply. Access covered the primary abstract and selected replication text, plus the reply's argument; no complete numerical audit of Folkes's studies is claimed. A reply is not an independent replication.
12. Petty, Cacioppo, and Goldman (1981), Personal involvement as a determinant of argument-based persuasion. The complete main article was read and Table 1 visually checked. The factorial experiment had 145 participants, with a separate 18-person no-message baseline. The result concerns attitudes, not behavior. Book [06.0036]–[06.0038].
13. Fillon et al. (2023), The Effectiveness of the 'But-you-are-free' technique. Methods, results, risk-of-bias analysis, and discussion were inspected. Overall  $g=.44$ ; seven low-risk studies  $g=.11$ , 95% CI  $-.18$  to  $.40$ . A low-risk-plus-some-concerns subset remained positive. These are standardized differences, not percentage uplifts. The source was published after this edition. Book [11.0111]–[11.0113], [17.0112].
14. Meeker et al. (2014), Nudging Guideline-Concordant Antibiotic Prescribing. Final methods, results, and all four tables were inspected; Tables 3–4 were checked visually. Table 4 gives 43.5%→33.7% for poster clinicians and 42.8%→52.7% for controls; the abstract reverses the two baseline labels. The controlled difference in changes is  $-19.7$  percentage points, 95% CI  $-33.04$  to  $-5.8$ . The appendix's calculation distinguishes both sets of labels, conventional relative changes, and the unverified origin of the book's 27%. Book [12.0230].
15. Sallis et al. (2020), Prescriber Commitment Posters to Increase Prudent Antibiotic Prescribing in English General Practice. Main methods, tables, results, and limitations were read. The 196 randomized units included usual care, posters, and posters plus an automated telephone message. The main poster outcome was overall dispensing; favorable secondary message analyses do not establish a primary poster benefit. The study predates this edition.
16. Book [12.0216]–[12.0224], especially [12.0220] and [12.0224]; public-goal counterexample [17.0127]. These support Cialdini's tentative Iowa explanation. The exact underlying Iowa quantities remain unresolved. The appendix separately examines the book's interpretation of the public-goal counterexample.
17. Unity's scope and qualifications: [13.0010]–[13.0013], [13.0026]–[13.0030], [13.0126]–[13.0194], [13.0236]–[13.0251], [17.0165], [17.0182]. The biological-source relationships are examined separately in appendix U02. The self–other-overlap illustration is a scale, not evidence of neural merger.
18. Mousa (2020), Building social cohesion between Christians and Muslims through soccer in post-ISIS Iraq. Main design, results, tables/figures, and conclusion were read. Staff contacted players at the six-month follow-up to record training with Muslims; this was not six months of continuously observed conduct. The book cites the study in a general caution about contact [17.0182], rather than spelling out its soccer-specific boundary.
19. Greenwood et al. (2020), Physician–patient racial concordance and disparities in birthing mortality for newborns, and Borjas and VerBruggen (2024), Physician–patient racial concordance and newborn mortality. Both main articles and their relevant tables were read; supplements and participant-level analyses were not rerun. The original paper acknowledges selection, omitted variables, and unknown mechanisms. The later analysis changes case-mix specifications using the same administrative records. This is not an independent-population replication or a demonstration that racial matching has exactly no effect. Book [17.0150]. Appendix U06 distinguishes the interaction from a within-race physician comparison.
20. Sagarin et al. (2002), Dispelling the illusion of invulnerability. All three experiments' methods/results were inspected. The delayed test followed 55 of 130 initial participants after 1–4 days. The 2004 correction restores omitted response-scale labels; it does not withdraw findings. Book [10.0129]–[10.0130], [17.0105].

21. Ott et al. (2011), Finding Deceptive Opinion Spam by Any Stretch of the Imagination. Dataset methods, human judgments, Tables 2–3, and validation were inspected. The 800 positive hotel reviews were half commissioned fakes and half selected reviews presumed genuine, not independently authenticated stays. The human subset contained 160 reviews. The three human judges' accuracies were 61.9%, 56.9%, and 53.1%; the essay's 53–62% range rounds those values. Book headline [06.0056], computer explanation [06.0057], and skills caveat [06.0065].
22. Gift-redefinition language [07.0164], [07.0169], and more qualified summary [07.0179]. Scarcity-awareness limits [11.0169]–[11.0172]. Distinguishing a reason to refuse from the disappearance of a feeling is editorial reasoning, not a tested intervention effect attributed to the book.
23. The book's threshold is a proven major violation or multiple proven minor violations [13.0251]; also [13.0241]–[13.0242], [17.0185]–[17.0187]. Kouchaki, Gino, and Feldman (2019) and Cialdini and colleagues' Bad Apples paper were accessible at abstract/preview scope. Their reported outcomes do not establish a comparative test of dismissal policies. No misconduct allegation or measured harm rate is inferred.